



Article

How Perceived Fit Shapes Value Co-Creation and Co-Destruction in Intelligent Customer Service: Psychological Mechanisms and Moderation by Digital Self-Efficacy

Lei Wang, Jiayi Ren *, Shiyi Sun and Yitao Chen 

School of business, Guilin University of Electronic Technology, Guilin 541004, China; wangleiss@guet.edu.cn (L.W.); 15257435765@163.com (S.S.); ace_chen88@guet.edu.cn (Y.C.)

* Correspondence: renjiayigu@163.com

Abstract

Intelligent customer service is increasingly used to reduce service costs and improve efficiency, yet users often experience divergent outcomes, ranging from value co-creation to value co-destruction. Using task–technology fit theory as the primary theoretical lens, this study examines how the alignment between users’ service task requirements and intelligent customer service capabilities is associated with divergent value outcomes. Within this framework, cognitive load and perceived control are theorized as two localized psychological mechanisms, human–AI trust as a conversion mechanism, and digital self-efficacy as a boundary condition. Using two scenario-based experiments, we develop and test the proposed model. Study 1 shows that participants in the high-fit condition reported higher value co-creation and lower value co-destruction than those in the low-fit condition. The results further reveal an asymmetric mechanism: cognitive load did not directly explain value co-creation but was more strongly associated with value co-destruction, whereas perceived control showed a broader association with both value outcomes. Study 2 extends Study 1 by examining digital self-efficacy as a boundary condition and shows that the conditional indirect association patterns vary across levels of digital self-efficacy. These findings contribute to extending task–technology fit theory to post-use value formation and provide cautious implications for intelligent customer service design and management.

Keywords: intelligent customer service; value co-creation; value co-destruction; task–technology fit theory; digital self-efficacy



Academic Editor: Ja-Shen Chen

Received: 15 May 2026

Revised: 9 July 2026

Accepted: 9 July 2026

Published: 11 July 2026

Copyright: © 2026 by the authors.

Licensee MDPI, Basel, Switzerland.

This article is an open access article distributed under the terms and conditions of the [Creative Commons Attribution \(CC BY\)](https://creativecommons.org/licenses/by/4.0/) license.

1. Introduction

Intelligent customer service refers to AI-based service robots that are capable of interacting with users in natural language, autonomously understanding user needs, and performing service tasks [1]. With the rapid development of artificial intelligence and natural language processing technologies, intelligent customer service is increasingly replacing human agents and has become an important means for firms to reduce service costs, improve operational efficiency, and respond to customers at scale. In contemporary digital service environments, intelligent customer service is no longer a supplementary tool but has become a key interface through which users interact with firms. It is playing an increasingly visible role in business operations and everyday digital service encounters [2], while also emerging as a practical tool for cost reduction and service efficiency. According to the 2025 New-Generation Artificial Intelligence Customer Service Market Research

Report released by Hongxiong AI, China's intelligent customer service market is expected to exceed RMB 60 billion in 2025, and this market is projected to continue expanding in the coming years. This rapid expansion makes it increasingly important to understand not only whether users adopt intelligent customer service, but also how they experience it and what value outcomes arise during actual interactions.

However, although intelligent customer service offers significant advantages in reducing costs and improving efficiency, its service outcomes are not always satisfactory. Different users may generate markedly different value outcomes. Some users are satisfied with the efficient information delivery and decision support of intelligent customer service, continue to use it, and co-create value with firms [3]. Others, by contrast, may feel frustrated due to problems such as AI hallucinations, poor information quality, and a lack of empathy [4,5], and may even abandon the service, ultimately resulting in value co-destruction. This phenomenon of "the same technology but different perceptions" highlights the core issue of user experience divergence in intelligent customer service. Such divergence indicates that intelligent customer service may generate not only positive experiences but also negative responses during user interactions. Therefore, the practical challenge for firms is not simply how to deploy intelligent customer service, but how to ensure that such systems create, rather than destroy, value in user interactions.

These practical concerns are also echoed in the service literature. In practice, AI-enabled frontline technologies have been increasingly deployed to transform service encounters and improve service efficiency [6]. However, existing research suggests that scholarly understanding of AI-driven frontline service encounters remains fragmented and lacks sufficient coherence, leaving important questions about how such technologies should be effectively implemented in service interactions [7]. This limitation is especially salient in intelligent customer service contexts, where AI-enabled service interactions may generate value co-creation, value co-destruction, or both, particularly when AI systems fail to meet users' needs or expectations [8]. Therefore, firms still need clearer guidance on how to leverage intelligent customer service to enhance user experience, promote value co-creation, and prevent value co-destruction. From a theoretical perspective, existing studies have largely focused on the antecedents of intelligent customer service adoption, such as the effects of perceived usefulness and perceived ease of use on adoption intention [9], or on how the technological characteristics of intelligent customer service influence user satisfaction [10]. In contrast, insufficient attention has been paid to the dynamic psychological processes users experience during actual interactions with intelligent customer service and to how these processes may lead to divergent value outcomes. Moreover, prior research has largely overlooked the boundary role of individual digital self-efficacy in shaping these processes [11].

To address these gaps, this study uses task-technology fit theory as the primary theoretical lens to explain why intelligent customer service may generate divergent value outcomes. Perceived fit captures the extent to which the system's response capability aligns with users' service task requirements during human-AI interactions. Within this fit-based framework, cognitive load and perceived control are conceptualized as two localized psychological mechanisms: cognitive load reflects the information-processing cost imposed by the interaction, whereas perceived control reflects users' sense of agency and controllability during the service process. Human-AI trust is further theorized as a conversion mechanism through which these psychological experiences are linked to value co-creation or value co-destruction. In addition, digital self-efficacy is introduced as a boundary condition that explains why users may respond differently to the same intelligent customer service condition.

On this basis, this study makes several theoretical and practical contributions. Theoretically, first, it extends task–technology fit theory to the post-use value formation of intelligent customer service by explaining how fit or misfit between user tasks and system capabilities may be associated with divergent value outcomes. Second, it moves beyond the dominant focus of prior research on adoption intention and technological attributes by examining the psychological mechanisms that unfold during actual human–AI service interactions. Specifically, cognitive load and perceived control are positioned as two localized mechanisms within the task–technology fit framework, while human–AI trust is treated as a conversion mechanism that links users’ interaction experiences to value co-creation or value co-destruction. Third, this study identifies digital self-efficacy as a boundary condition that helps explain why users may respond differently to the same intelligent customer service condition. Practically, this study suggests that firms should not only improve the functional fit between intelligent customer service and users’ task requirements, but also reduce unnecessary cognitive burden, preserve users’ perceived control, and foster appropriate human–AI trust.

2. Theoretical Background and Hypothesis Development

2.1. Intelligent Customer Service, Value Co-Creation, and Value Co-Destruction

Service-dominant logic (S-D logic) distinguishes between two possible value outcomes arising from interactions between users and firms: value co-creation and value co-destruction [12,13]. In this study, these two constructs are examined as user-level value outcomes reflected in post-interaction evaluations and behavioral intentions, rather than directly observed long-term value behaviors. Value co-creation (VCC) refers to the joint creation of value through user–firm interactions and is reflected in favorable post-interaction responses, such as user satisfaction, continued usage intention, and favorable word of mouth [14]. By contrast, value co-destruction (VCD) refers to value diminution caused by resource misuse or interaction failure and is manifested in negative post-interaction responses, such as user frustration, complaints, negative word of mouth, and service discontinuance [15]. Together, these two value outcomes constitute the spectrum of value formation in users’ interactions with intelligent customer service.

Although value co-creation and value co-destruction are both possible outcomes of service interactions, prior research cautions against treating value co-destruction simply as the reverse side of value co-creation. Instead, value co-destruction should be understood as a distinct yet dynamically related value-formation process that may emerge through negative or insufficient service interaction experiences, such as service failure, emotional friction, accumulated frustration, and active resistance during user-system interactions [16,17]. In intelligent customer service contexts, this distinction is particularly important because AI-enabled service encounters may simultaneously provide efficiency and convenience while also generating misunderstanding, loss of control, distrust, and frustration. Therefore, this study conceptualizes value co-creation and value co-destruction as two divergent value trajectories that may be shaped by different psychological mechanisms rather than as two symmetric endpoints on a single continuum.

Task–technology fit theory provides a useful theoretical lens for understanding value outcomes in the use of intelligent customer service [18]. This theory posits that the extent to which a technology can effectively support users in accomplishing service tasks depends on the degree of fit between its functionalities and users’ task requirements. A high level of fit generally leads to greater interaction efficiency and better task performance, whereas a low level of fit is more likely to create use barriers and undermine service effectiveness [19,20]. When the functionalities of intelligent customer service are highly aligned with users’ task requirements, that is, when users perceive a high level of fit, they are more likely to

believe that the system can effectively support task completion and thus enjoy a more efficient and seamless interaction experience. On this basis, users are also more likely to develop positive interactions with the system and, through continued use, facilitate value co-creation [21]. Conversely, when the degree of fit is low, intelligent customer service may fail to understand user intent, involve complex and disorganized interaction processes, or fail to complete the task successfully. Under such circumstances, users are more likely to experience negative emotions, evaluate the system less favorably, and respond with discontinuance, complaints, or other negative reactions, thereby increasing the likelihood of value co-destruction through failed interaction processes [22]. Differences in task–technology fit thus constitute the logical starting point for the divergence of users' value outcomes.

Integrating task–technology fit theory with the value formation perspective, perceived fit can be regarded as a key starting condition for explaining the divergence in users' value outcomes. Higher perceived fit is expected to be associated with stronger value co-creation and lower value co-destruction, whereas lower perceived fit is more likely to be associated with value co-destruction. Accordingly, the following hypotheses are proposed:

H1. *Perceived fit is significantly associated with value co-creation and value co-destruction.*

H1a. *Perceived fit is positively associated with value co-creation.*

H1b. *Perceived fit is negatively associated with value co-destruction.*

2.2. Perceived Fit as the Primary Theoretical Anchor: The Efficiency–Autonomy Tension

Perceived fit serves as the primary theoretical anchor of this study. From the task–technology fit perspective, fit does not simply mean that a technology is useful or easy to use; rather, it captures whether the system's capabilities match the specific requirements of users' service tasks [18–20]. In intelligent customer service encounters, this alignment is critical because users rely on AI systems not only to obtain information but also to solve problems such as returns, refunds, complaints, and service recovery. When the system accurately understands users' intentions and provides task-relevant responses, users are more likely to experience the interaction as efficient, understandable, and manageable. When the system misunderstands the task or traps users in repetitive interaction loops, the same automated interface may increase cognitive burden, reduce perceived agency, and produce negative value outcomes.

Importantly, the fit–value relationship should not be understood as a frictionless positive sequence. AI-enabled service systems may reduce cognitive load by simplifying procedures and automating responses, but excessive automation may also weaken users' sense of control when the system restricts correction, offers limited explanations, or delays access to human agents. Thus, the psychological process linking perceived fit to value outcomes involves an efficiency–autonomy tension. This tension provides the basis for examining cognitive load and perceived control as two localized psychological mechanisms, and human–AI trust as the conversion mechanism through which these experiences are linked to value co-creation or value co-destruction.

2.3. The Mediating Role of Cognitive Load

Cognitive load is defined as the mental effort required for individuals to process and respond to information [23]. Cognitive load theory suggests that the working memory resources individuals rely on when processing new information are limited in both capacity and duration [24]. In intelligent customer service, however, reduced cognitive load should not be understood as an automatically positive experience. Cognitive ease and

perceived control are analytically distinct but practically intertwined. A highly automated system may reduce the effort required to complete a task, but if users cannot understand, interrupt, correct, or override the system, reduced cognitive load may coexist with a weakened sense of control. Therefore, cognitive load is treated as one localized mechanism within the task–technology fit framework, rather than as a self-sufficient explanation of value outcomes.

When the functionalities of intelligent customer service are highly aligned with users' task requirements, the interaction process tends to be smooth and the information presented clear, allowing users to complete tasks without expending excessive cognitive effort. Under such circumstances, cognitive load remains at a relatively low level, enabling users to allocate more cognitive resources to the task itself, which may support greater engagement, better task performance, and more satisfactory experiences [25–27], thereby facilitating value co-creation. By contrast, when the level of fit is low, users need to invest more mental resources to understand system feedback and complete the task, which may weaken their motivation and engagement and even induce negative reactions such as frustration and disengagement [25]. As a result, value co-destruction and other negative behavioral outcomes are more likely to occur. Based on this, the following hypotheses are proposed:

H2. *Cognitive load mediates the relationship between perceived fit and value co-creation/value co-destruction.*

H2a. *Perceived fit is positively associated with value co-creation through lower cognitive load.*

H2b. *Perceived fit is negatively associated with value co-destruction through lower cognitive load.*

2.4. The Mediating Role of Perceived Control

In human–AI interaction contexts, perceived control generally refers to individuals' experience of control over their own actions and their outcomes during interactions with technology, that is, the sense of agency subjectively perceived by users [28]. Perceived control is one of the basic psychological needs of human beings [29]. When this need is satisfied during the interaction process, individuals are more likely to develop intrinsic motivation and positive emotional experiences, thereby enhancing their satisfaction with and trust in the system and fostering positive behavioral intentions—that is, value co-creation, such as favorable evaluations of intelligent customer service, willingness to recommend it, and continued usage intention. In contrast, when perceived control is threatened or deprived, users may experience frustration, anxiety, or even resistance, which in turn may induce negative behavioral intentions—namely, value co-destruction—such as complaints, negative word of mouth, avoidance, or even discontinuance of use. In other words, perceived control helps explain how users evaluate the agency implications of intelligent customer service. High perceived control may support constructive engagement when users feel that the automated process remains understandable, influenceable, and reversible. Low perceived control, by contrast, may increase the likelihood of withdrawal, complaint, or resistance when users feel constrained by the system. Thus, perceived control is not simply the opposite of cognitive load; rather, it captures a distinct agency-based experience that may either complement or conflict with efficiency-oriented automation. Based on this agency-based logic, the following hypotheses are proposed:

H3. *Perceived control mediates the relationship between perceived fit and value co-creation/value co-destruction.*

H3a. *Perceived fit is positively associated with value co-creation through higher perceived control.*

H3b. *Perceived fit is negatively associated with value co-destruction through higher perceived control.*

2.5. The Chain Mediating Role of Human–AI Trust

According to automation trust theory, users' trust in a system is built on their evaluation of its performance. However, human–AI trust should not be understood as a stable or purely linear mechanism. Recent research on appropriate reliance and trust dynamics suggests that users may either over-rely on AI systems when trust exceeds system reliability or insufficiently rely on AI recommendations when trust falls below system capability [30,31]. Such trust miscalibration may lead to two different responses. On the one hand, automation bias reflects users' excessive reliance on AI-generated outputs; on the other hand, algorithm aversion reflects users' tendency to avoid or resist algorithmic recommendations after experiencing errors [32,33]. These insights are particularly relevant to intelligent customer service because perceived fit may encourage appropriate reliance by signaling system competence, whereas lower perceived fit may be associated with distrust, frustration, and resistance. Therefore, human–AI trust is conceptualized in this study not as a simple positive mediator, but as a conversion mechanism through which cognitive load and perceived control are translated into either value co-creation or value co-destruction. When a system helps users accomplish tasks efficiently at a low cognitive cost, users tend to attribute such a smooth experience to the system's professionalism and reliability, thereby enhancing human–AI trust [34]. Similarly, users' positive perceptions of the system's fluency, controllability, and competence during the interaction process may be translated into favorable judgments of its reliability and dependability, which in turn strengthen human–AI trust [35]. Building on automation trust theory, users may make a "performance–effort" trade-off when deciding how much system assistance to accept [36]. In intelligent customer service, however, accepting system assistance does not necessarily mean relinquishing perceived control. When the system's support is transparent, task-relevant, and correctable, users may allow the system to handle routine procedures while still maintaining a sense of agency over the interaction. In this case, perceived control is not weakened; rather, users experience the system's assistance as controllable and competence-enhancing. Therefore, perceived control may strengthen human–AI trust by allowing users to feel that the system is both capable and manageable.

Accordingly, human–AI trust is treated as a conversion mechanism rather than a purely positive mediator. It represents a psychological link through which users' experiences of efficiency and agency are associated with different forms of post-interaction response. When users interpret low cognitive burden and retained control as signals of system competence and reliability, trust may support constructive engagement and value co-creation. When cognitive effort increases, control is weakened, or trust becomes miscalibrated, users may withdraw from the interaction, complain, resist the system, or discontinue use, thereby increasing the likelihood of value co-destruction. This logic does not assume a frictionless chain from fit to positive outcomes; instead, it emphasizes that trust formation depends on how users interpret the balance between automation efficiency and perceived agency. Based on the above analysis, this study proposes the following hypotheses:

H4. *Cognitive load and human–AI trust jointly play a chain mediating role in the relationship between perceived fit and value outcomes.*

H4a. *Perceived fit in intelligent customer service is positively associated with value co-creation through lower cognitive load and stronger human–AI trust.*

H4b. *Perceived fit in intelligent customer service is negatively associated with value co-destruction through lower cognitive load and stronger human–AI trust.*

H5. *Perceived control and human–AI trust jointly play a chain mediating role in the relationship between perceived fit and value outcomes.*

H5a. *Perceived fit in intelligent customer service is positively associated with value co-creation through higher perceived control and stronger human–AI trust.*

H5b. *Perceived fit in intelligent customer service is negatively associated with value co-destruction through higher perceived control and stronger human–AI trust.*

2.6. The Moderating Role of Digital Self-Efficacy

Digital self-efficacy (DSE), rooted in self-efficacy theory, refers to individuals' confidence in and perceived capability to use digital technologies to accomplish tasks [11]. In interactions with intelligent customer service, users with high DSE possess stronger information-processing abilities and greater confidence in coping with technology, enabling them to handle system feedback more efficiently and adapt more quickly to the pace of interaction. Therefore, for users with high DSE, the negative association between perceived fit and cognitive load is expected to be more pronounced. Even when the level of perceived fit is low, they are better able to mobilize cognitive resources to cope with interaction barriers, thereby mitigating the negative impact of cognitive burden. Accordingly, the following hypothesis is proposed:

H6. *Digital self-efficacy moderates the association between perceived fit and cognitive load. Specifically, for users with high DSE, the negative association between perceived fit and cognitive load in intelligent customer service is more pronounced.*

Control-related changes and their effects on trust are more complex in intelligent customer service. In this study, perceived control does not mean that users must manually control every step of the service process. Rather, it refers to users' subjective sense that the interaction remains understandable, influenceable, and reversible. In AI-enabled service interactions, users may allow the system to handle routine operations while still feeling in control if the process is transparent, task-relevant, and allows correction, interruption, or transfer to a human agent when necessary. Therefore, controllable assistance should not be interpreted as a loss of perceived control, but as a form of control experience that emerges when users voluntarily rely on a competent system while retaining perceived agency over the interaction.

Users with high digital self-efficacy are more capable of evaluating whether system agency can help them complete the task efficiently. When the intelligent customer service system fits their task requirements, high-DSE users are more likely to interpret system assistance as controllable and competence-enhancing, thereby reporting stronger perceived control and greater human–AI trust. By contrast, users with low DSE may be less able to evaluate or manage system agency and may therefore experience weaker perceived control even under similar service conditions. Accordingly, the following hypothesis is proposed:

H7. *Digital self-efficacy moderates the association between perceived fit and perceived control. Specifically, for users with high DSE, the positive association between perceived fit and perceived control in intelligent customer service is more pronounced.*

These moderation patterns are further expected to be reflected in the conditional indirect association paths. Along the cognitive load path, users with high DSE are better able to translate reductions in cognitive load into recognition of the system's competence and greater trust, thereby being associated with stronger value co-creation and lower value

co-destruction. Along the perceived control path, users with high DSE are more likely to interpret the enhancement of perceived control as a signal of the system's reliability, thereby being associated with stronger conditional indirect association patterns along the “perceived fit → perceived control → human–AI trust → value co-creation/value co-destruction” path. Accordingly, the following hypotheses are proposed:

H8. *Digital self-efficacy moderates the chain path of “perceived fit → cognitive load → human–AI trust → value co-creation/value co-destruction.”*

H9. *Digital self-efficacy moderates the chain path of “perceived fit → perceived control → human–AI trust → value co-creation/value co-destruction.”*

Based on the above hypotheses, this study proposes the following theoretical framework, as shown in Figure 1.

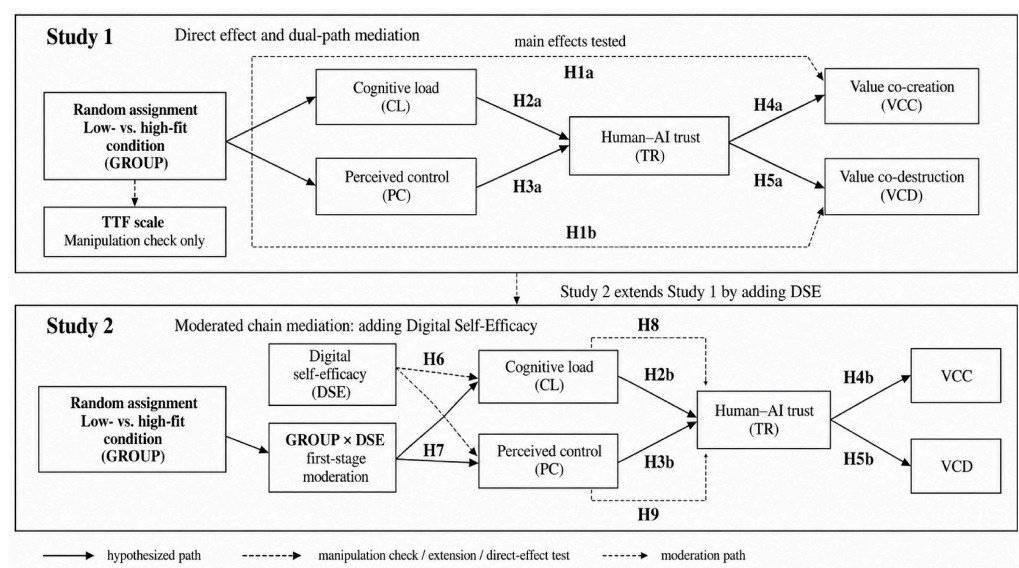


Figure 1. Research model.

3. Study 1: The Dual-Path Mechanism Linking Perceived Fit to Users' Value Outcomes

Study 1 examines whether experimentally manipulated fit condition is associated with systematic differences in users' post-scenario psychological evaluations, including cognitive load, perceived control, and human–AI trust, as well as value-related behavioral intentions, including value co-creation and value co-destruction. It also examines whether these associations are consistent with the theoretically proposed mediating and serial mediating roles of cognitive load, perceived control, and human–AI trust.

3.1. Pilot Study

3.1.1. Participants and Procedure

The purpose of the pilot study was to examine whether the scenario manipulation of the core independent variable, perceived fit, operationalized through two randomly assigned experimental conditions (1 = high-fit condition, 0 = low-fit condition), was successful. A one-way analysis of variance (ANOVA) was conducted, with experimental condition as the independent variable.

The pilot study questionnaire was distributed and collected through the online survey platform Wenjuanxing. First, all participants were randomly assigned to either the high-fit condition or the low-fit condition. They were then guided to imagine the following scenario

of interacting with intelligent customer service: “You purchased a coat in a livestreaming room on an e-commerce platform. After receiving the product, you found that the coat had obvious quality problems. You would like to apply for a return or exchange and seek help through the customer service portal in the livestreaming room.” Next, participants in each group were asked to read the corresponding scenario materials. After reading the scenario, they were required to complete the task–technology fit scale (TTF) based on their actual feelings. The scale consisted of four items adapted from the study of Goodhue and Thompson [37], with a sample item being: “I feel that the functions of this intelligent customer service system can meet my needs”. All items were measured on a seven-point Likert scale (1 = strongly disagree, 7 = strongly agree). Finally, participants provided basic demographic information. A total of 162 valid questionnaires were collected in the pilot study, including 81 participants in the low-fit condition and 81 participants in the high-fit condition.

3.1.2. Pilot Study Results

Descriptive statistics showed (Figure 2) that participants in the high-fit condition reported a significantly higher score on task–technology fit ($M = 5.03$, $SD = 1.32$) than those in the low-fit condition ($M = 3.04$, $SD = 1.40$). The ANOVA results further indicated a highly significant difference between the two groups in task–technology fit, $F(1, 160) = 86.41$, $p < 0.001$, $\eta^2 = 0.35$. According to Cohen’s (1988) criteria, this effect size can be classified as large ($\eta^2 > 0.14$), suggesting that the scenario materials were highly effective in manipulating the independent variable, perceived fit [38]. In addition, the test of homogeneity of variance showed that Levene’s test was not significant ($p = 0.617 > 0.05$), indicating that the assumption of homogeneity of variance was satisfied and further supporting the reliability of the analysis results.

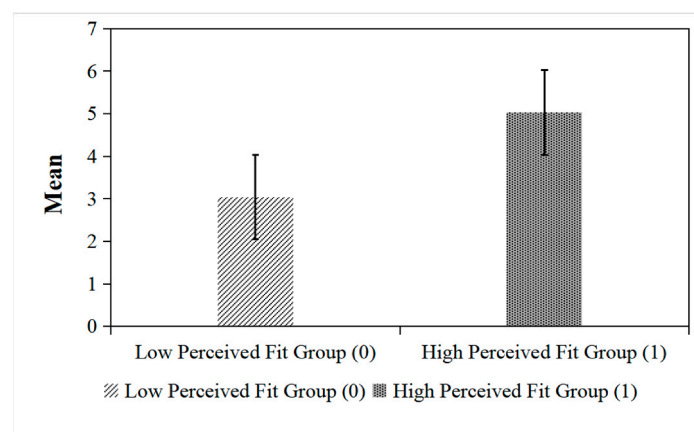


Figure 2. One-way ANOVA of perceived fit.

In summary, the pilot study results indicate that the scenario materials developed in this study were effective in manipulating participants’ perceived fit and can, therefore, be used in the formal experiment.

3.2. Main Study

3.2.1. Study 1 Experimental Design and Procedure

The target population of the main study consisted of individuals who had prior experience using intelligent customer service. Before the formal data collection, the minimum required sample size was estimated using G*Power 3.1. An independent-samples *t*-test was selected (large effect size $d = 0.8$, $\alpha = 0.05$, statistical power = 0.80), and the minimum total sample size was calculated to be 142. Online samples were then collected through the

Wenjuanxing platform, resulting in 256 valid questionnaires and an effective response rate of 96.60%. In terms of gender, 45.7% of the participants were male and 54.3% were female. In terms of the frequency of intelligent customer service use, 27.0% of the participants had used intelligent customer service 1–5 times in the past six months, 43.4% had used it 6–10 times, and 29.7% had used it more than 10 times. The study procedure was as follows.

The main study was also conducted through the Wenjuanxing platform for questionnaire distribution and collection. First, all participants were randomly assigned to either the high-fit condition ($n = 135$) or low-fit condition ($n = 121$). It should be noted that perceived fit in this study was operationalized as a randomly assigned experimental condition rather than as a post-vignette self-reported predictor. Specifically, the scenarios manipulated the objective alignment between users' service task requirements and the intelligent customer service system's response capability. In the high-fit condition, the system accurately understood the user's request and provided effective task support, whereas in the low-fit condition, the system misunderstood the user's intention and increased the complexity of task completion. The post-vignette task–technology fit measure was used only as a manipulation check to verify whether the scenario successfully induced different levels of perceived fit. Thus, the focal independent variable in the structural model was the randomly assigned experimental condition, not participants' post-vignette subjective perceived fit ratings. This design helps reduce concerns associated with using a post-scenario subjective perception as the main predictor and strengthens the interpretation that the experimental condition shaped participants' subsequent psychological evaluations and behavioral intentions.

Participants were then asked to imagine the following scenario of interacting with intelligent customer service: "You purchased a coat in a livestreaming room on an e-commerce platform. After receiving the product, you found that the coat had obvious quality problems. You would like to apply for a return or exchange and seek help through the customer service portal in the livestreaming room". Next, participants in each group were asked to read the corresponding scenario materials. After that, they were required to complete a questionnaire based on their actual feelings, including measures of cognitive load ($\alpha = 0.973$), perceived control ($\alpha = 0.952$), human–AI trust ($\alpha = 0.970$), value co-creation ($\alpha = 0.905$), and value co-destruction ($\alpha = 0.883$). Specifically, cognitive load was measured with six items adapted from Hong et al. [39]; perceived control was measured with four items adapted from the perceived control scale developed by Collier and Barnes (2015) [40], with contextual revisions based on recent chatbot applications such as Lubbe (2025) [41]; the five items measuring human–AI trust were contextually adapted from the scales used by Hu et al. (2010), Alagarsamy and Mehroli (2023), and other scholars [42,43]; the four items measuring value co-creation were adapted from the work of Yang (2018) and other scholars [44]; and value co-destruction was measured with four items based on the work of Vargo and Lusch (2008) and other scholars [45]. All items in this study were measured on a seven-point Likert scale (1 = strongly disagree, 7 = strongly agree). Finally, participants provided their basic demographic information.

3.2.2. Study 1 Results

1. Measurement Model, Reliability, and Validity

A confirmatory factor analysis was conducted to assess the measurement model (Table 1). The six-factor model included task–technology fit, cognitive load, perceived control, human–AI trust, value co-creation, and value co-destruction. The model showed an acceptable fit to the data: $\chi^2(309) = 721.864$, CFI = 0.963, TLI = 0.958, RMSEA = 0.072, and SRMR = 0.022. All standardized factor loadings were statistically significant, ranging from 0.634 to 0.966. Composite reliability (CR) values ranged from 0.884 to 0.978, exceeding

the recommended threshold of 0.70. Average variance extracted (AVE) values ranged from 0.660 to 0.919, exceeding the recommended threshold of 0.50, thereby supporting convergent validity.

Table 1. Reliability and convergent validity of the constructs in Study 1.

Construct	Items	Standardized Factor Loadings	CR	AVE
Task–technology fit (TTF)	TTF1–TTF4	0.936–0.966	0.978	0.919
Cognitive load (CL)	CL1–CL6	0.888–0.956	0.973	0.857
Perceived control (PC)	PC1–PC4	0.893–0.934	0.952	0.835
Human–AI trust (TR)	TR1–TR5	0.900–0.953	0.970	0.865
Value co-creation (VCC)	VCC1–VCC4	0.686–0.944	0.905	0.717
Value co-destruction (VCD)	VCD1–VCD4	0.634–0.910	0.883	0.660

2. Manipulation Check

To ensure the effectiveness of the manipulation of the core independent variable, perceived fit, this study used task–technology fit as the key criterion variable and conducted a one-way analysis of variance (ANOVA) between the high-fit condition (GROUP = 1) and the low-fit condition (GROUP = 0). The results showed that participants in the high-fit condition scored significantly higher on task–technology fit ($M = 5.74$, $SD = 0.75$) than those in the low-fit condition ($M = 2.03$, $SD = 0.84$), with a significant between-group difference, $F(1, 254) = 1387.132$, $p < 0.001$, $\eta^2 = 0.845$. The mean difference between the two groups was 3.71. These results indicate that the two intelligent customer service scenarios designed in this study successfully induced differences in users' core perceptions of task–technology fit, confirming that the manipulation of perceived fit was effective.

3. Main Effects

Using perceived fit in intelligent customer service, operationalized as the randomly assigned experimental condition (1 = high-fit condition, 0 = low-fit condition), as the independent variable, one-way analyses of variance (ANOVAs) were conducted separately with value co-creation and value co-destruction as the dependent variables. The results showed that perceived fit had a significant main effect on value co-creation, $F(1, 254) = 503.25$, $p < 0.001$, partial $\eta^2 = 0.665$, and also a significant main effect on value co-destruction, $F(1, 254) = 494.88$, $p < 0.001$, partial $\eta^2 = 0.661$. Further analyses showed that participants in the high-fit condition reported a significantly higher level of value co-creation ($M = 5.51$, $SD = 0.79$) than those in the low-fit condition ($M = 2.89$, $SD = 1.07$), with a difference of 2.619 points, $p < 0.001$. In contrast, participants in the high-fit condition reported a significantly lower level of value co-destruction ($M = 2.47$, $SD = 0.89$) than those in the low-fit condition ($M = 4.96$, $SD = 0.91$), with a difference of 2.496 points, $p < 0.001$.

These results provide evidence that the experimentally manipulated fit condition was associated with systematic differences in users' value-related responses. Compared with participants in the low-fit condition, those in the high-fit condition reported higher value co-creation and lower value co-destruction. Therefore, H1, H1a, and H1b were supported.

4. Mediation Analysis

To examine the indirect association patterns through which perceived fit was linked to value co-creation and value co-destruction, mediation analyses were conducted using PROCESS Macro (Model 80; Hayes, 2022) [46]. The number of bootstrap resamples was set at 5000, and 95% confidence intervals for the indirect effects were calculated. The results are presented in Tables 2 and 3.

Table 2. Mediation analysis of the effect of perceived fit in intelligent customer service on value co-creation ($n = 256$).

Path	Effect	Boot SE	95% Bootstrap CI
Total indirect effect	3.1016	0.2570	[2.6067, 3.6148]
Direct effect	−0.4824	0.2628	[−0.9999, 0.0352]
Perceived fit → Cognitive load → Value co-creation	−0.0976	0.3298	[−0.6889, 0.5973]
Perceived fit → Perceived control → Value co-creation	1.0201	0.2738	[0.4585, 1.5458]
Perceived fit → Human–AI trust → Value co-creation	0.5220	0.1717	[0.1839, 0.8601]
Perceived fit → Cognitive load → Human–AI trust → Value co-creation	1.2671	0.2754	[0.7749, 1.8680]
Perceived fit → Perceived control → Human–AI trust → Value co-creation	0.3900	0.1517	[0.0793, 0.6779]

Table 3. Mediation analysis of the effect of perceived fit in intelligent customer service on value co-destruction ($n = 256$).

Path	Effect	Boot SE	95% Bootstrap CI
Total indirect effect	−2.5994	0.3032	[−3.2523, −2.0593]
Direct effect	0.1031	0.3196	[−0.5265, 0.7326]
Perceived fit → Cognitive load → Value co-destruction	−1.4294	0.4005	[−2.2436, −0.6786]
Perceived fit → Perceived control → Value co-destruction	−0.6175	0.3311	[−1.3102, −0.0150]
Perceived fit → Human–AI trust → Value co-destruction	−0.1324	0.0925	[−0.3571, −0.0056]
Perceived fit → Cognitive load → Human–AI trust → Value co-destruction	−0.3213	0.1621	[−0.6656, −0.0333]
Perceived fit → Perceived control → Human–AI trust → Value co-destruction	−0.0989	0.0705	[−0.2686, −0.0003]

Using value co-creation as the dependent variable, the total effect of perceived fit on value co-creation was 2.619, with a 95% confidence interval (CI) of [2.456, 3.118]; the direct effect was −0.482, with a 95% CI of [−0.999, 0.035]; and the total indirect effect was 3.102, with a 95% CI of [2.607, 3.615]. Among these effects, the confidence intervals for both the total effect and the total indirect effect did not include zero, whereas the confidence interval for the direct effect included zero, indicating that the effect of perceived fit on value co-creation was fully mediated. Regarding the specific paths, among the simple mediation paths, the indirect effect of perceived fit on value co-creation through perceived control was significant (effect = 1.020, 95% CI = [0.459, 1.546]), and the indirect effect through human–AI trust was also significant (effect = 0.522, 95% CI = [0.184, 0.860]), whereas the indirect effect through cognitive load was not significant (95% CI = [−0.689, 0.597]). Among the chain mediation paths, the indirect effect of perceived fit on value co-creation through “cognitive load → human–AI trust” was significant (effect = 1.267, 95% CI = [0.775, 1.868]), and the indirect effect through “perceived control → human–AI trust” was also significant (effect = 0.390, 95% CI = [0.079, 0.678]). These findings suggest that, in the value co-creation path, cognitive load did not exhibit a significant simple mediating effect but mainly operated through the chain path of “cognitive load → human–AI trust”. This indicates that lower cognitive load alone is insufficient to automatically translate into positive value outcomes; only when lower cognitive load is further interpreted by users as a signal of the system’s competence, reliability, and dependability can it promote value co-creation through enhanced human–AI trust. By contrast, perceived control not only exerted a significant simple mediating effect, but also further influenced value co-creation through human–AI trust. This suggests that the formation of positive value outcomes driven by perceived fit in intelligent customer service is not directly promoted by a single psychological mechanism but rather depends on the further transformation from cognitive experience to trust judgment. Accordingly, H2a was not supported, whereas H3a, H4a, and H5a were supported. This finding indicates that cognitive load does not

translate into value co-creation in a simple or direct way but works mainly through users' trust judgments.

Using value co-destruction as the dependent variable, the total indirect effect of perceived fit on value co-destruction was -2.599 , with a 95% confidence interval (CI) of $[-3.252, -2.059]$, whereas the direct effect was 0.103 , with a 95% CI of $[-0.527, 0.733]$. The confidence interval for the total indirect effect did not include zero, whereas the confidence interval for the direct effect included zero, indicating that the effect of perceived fit on value co-destruction was fully mediated. Regarding the specific paths, among the simple mediation paths, the indirect effect of perceived fit on value co-destruction through cognitive load was significantly negative (effect = -1.429 , 95% CI = $[-2.244, -0.679]$), the indirect effect through perceived control was also significantly negative (effect = -0.618 , 95% CI = $[-1.310, -0.015]$), and the indirect effect through human–AI trust was likewise significantly negative (effect = -0.132 , 95% CI = $[-0.357, -0.006]$). Accordingly, H2b and H3b were supported. Among the chain mediation paths, the indirect effect of perceived fit on value co-destruction through “cognitive load $\rightarrow$ human–AI trust” was significantly negative (effect = -0.321 , 95% CI = $[-0.666, -0.033]$), and the indirect effect through “perceived control $\rightarrow$ human–AI trust” was also significantly negative (effect = -0.099 , 95% CI = $[-0.269, -0.0003]$). Therefore, H4b and H5b were supported.

3.2.3. Discussion of Study 1

Study 1 successfully manipulated perceived fit and showed that experimentally induced fit condition was associated with systematic differences in users' post-scenario psychological evaluations and value-related responses. Participants in the high-fit condition reported higher value co-creation and lower value co-destruction than those in the low-fit condition. Importantly, these results should not be interpreted as evidence that value co-creation and value co-destruction are simply reversible outcomes. Although the high-fit condition was associated with higher value co-creation and lower value co-destruction, the underlying paths were not identical. Cognitive load did not directly explain value co-creation but showed a stronger association with value co-destruction, whereas perceived control showed a more consistently constructive role across both outcomes. More importantly, Study 1 suggests that value divergence in intelligent customer service does not unfold through a single, symmetrical psychological process, but is linked to two distinct paths: cognitive load and the experience of control. The large effect sizes in Study 1 indicate that the scenario manipulation created a clear contrast between the high-fit and low-fit conditions. However, because cognitive load, perceived control, human–AI trust, value co-creation, and value co-destruction were measured after the same scenario exposure, the serial ordering among these variables should be interpreted as theoretically specified rather than temporally established. Thus, Study 1 provides evidence that the manipulated fit condition was associated with post-scenario psychological and value-related responses, while stronger claims about the temporal ordering of the mediators require future longitudinal or field-based evidence.

Specifically, along the cognitive load path, perceived fit did not directly support value co-creation merely through lower cognitive load; rather, the results were more consistent with the chain path of “cognitive load $\rightarrow$ human–AI trust”. In contrast, along the value co-destruction path, cognitive load was associated with value co-destruction both directly and indirectly through human–AI trust. Along the control experience path, perceived control was associated with both value co-creation and value co-destruction, and these associations were further linked to human–AI trust. These findings indicate that human–AI trust plays a critical transforming role in both paths.

Overall, Study 1 suggests that value co-creation and value co-destruction in intelligent customer service are not simply opposite outcomes of the same psychological mechanism, but rather a process of value divergence jointly shaped by cognitive load and the experience of control in an asymmetric manner. On this basis, Study 2 further introduces digital self-efficacy to examine the boundary conditions of this mechanism.

4. Study 2: The Contingent Mechanism Linking Perceived Fit to Users' Value Outcomes: The Moderating Role of Digital Self-Efficacy

4.1. Study 2 Experimental Design and Procedure

Study 2 aims to examine the moderating role of digital self-efficacy in the dual-path chain mediation model linking perceived fit to users' value outcomes through cognitive load/perceived control and human–AI trust, that is, to explore the boundary conditions of individual heterogeneity in this mechanism. This study employed a single-factor, two-level between-subjects design (high-fit condition vs. low-fit condition), with digital self-efficacy treated as a continuous moderating variable. The independent variable was therefore treated as an experimentally manipulated condition rather than as a post-vignette self-reported predictor. Study 2 adopted the same manipulation logic and scenario materials validated in Study 1 but did not include a separate task–technology fit manipulation-check scale. Accordingly, perceived fit was operationalized through the randomly assigned scenario condition, and the moderated mediation analyses were based on the experimental condition rather than on participants' subjective post-scenario fit ratings. Accordingly, all hypotheses in Study 2 are interpreted as associative relationships derived from scenario-based experimental evidence rather than as definitive causal effects or temporal process confirmations.

Study 2 was designed as an extension of Study 1 rather than as a fully independent replication. Specifically, Study 1 tested the basic dual-path mechanism linking fit condition to value outcomes through cognitive load, perceived control, and human–AI trust. Study 2 used the same validated scenario logic to hold the core fit manipulation constant while introducing digital self-efficacy as a pre-measured boundary condition. This design was suitable for examining whether individual differences in digital self-efficacy changed users' psychological responses to the same fit or misfit condition. Thus, the purpose of Study 2 was not to claim an independent replication of Study 1, but to examine the conditional nature of the proposed process model.

In terms of stimulus materials, Study 2 adopted the scenario materials that had been successfully validated in Study 1. Based on both the pilot study and the formal experiment in Study 1, the scenario materials that effectively induced perceptual differences among participants were selected: participants in the high-fit condition read a scenario in which the intelligent customer service system accurately understood users' needs and resolved problems smoothly, whereas participants in the low-fit condition read a scenario in which the intelligent customer service system misunderstood users' intentions and the interaction process was complex and confusing. The form of the scenario materials remained consistent with that used in Study 1, namely, written descriptions that guided participants to imagine a specific conversational interaction with intelligent customer service, so as to ensure comparability between the two experiments.

With regard to the experimental procedure, Study 2 generally followed the same procedure as Study 1. First, digital self-efficacy was measured before the scenario manipulation to reduce potential contamination from the experimental scenario ($\alpha = 0.874$). The digital self-efficacy scale was adapted from Kim et al. [11] and included three items, such as "I believe that I can easily learn how to use intelligent customer service." Participants were then randomly assigned to either the high-fit condition ($n = 163$) or the low-fit condition

($n = 165$) and were asked to read the corresponding scenario materials. After reading the scenario, participants completed the questionnaire based on their own feelings. The questionnaire included measures of cognitive load ($\alpha = 0.971$), perceived control ($\alpha = 0.794$), human–AI trust ($\alpha = 0.967$), value co-creation ($\alpha = 0.912$), and value co-destruction ($\alpha = 0.846$). All items were measured on a seven-point Likert scale (1 = strongly disagree, 7 = strongly agree). Finally, participants provided their basic demographic information, including gender, age, occupation, education level, and frequency of intelligent customer service use.

The main study targeted consumers who had prior experience using intelligent customer service. Considering the possibility of invalid questionnaires in the process of online data collection, this study collected a total of 350 questionnaires through the Wenjuanxing platform. After screening, 328 valid questionnaires were retained, yielding an effective response rate of 93.7%. In the sample, 40.85% of the participants were male and 59.15% were female. Participants were mainly between 26 and 45 years old, with 58.5% aged 26–35 and 25.0% aged 36–45. The sample was relatively well educated, with 76.2% holding a bachelor's degree and 6.7% holding a master's degree or above. In terms of the frequency of intelligent customer service use, 66.8% of the participants had used intelligent customer service six times or more in the past six months, indicating that the sample had a solid foundation of user experience with intelligent customer service.

4.2. Study 2 Results

4.2.1. Measurement Model, Reliability, and Validity

Before testing the hypothesized moderated mediation model, a confirmatory factor analysis was conducted to assess the measurement model (Table 4). The six-factor model included digital self-efficacy, cognitive load, perceived control, human–AI trust, value co-creation, and value co-destruction. The model fit was partially acceptable: $\chi^2(284) = 1368.602$, CFI = 0.905, TLI = 0.891, and SRMR = 0.096. All standardized factor loadings were statistically significant. Composite reliability (CR) values ranged from 0.850 to 0.960, exceeding the recommended threshold of 0.70. Average variance extracted (AVE) values ranged from 0.580 to 0.858, exceeding the recommended threshold of 0.50, thereby supporting convergent validity. To further assess potential common method variance and construct distinctiveness, we compared the proposed six-factor measurement model with a one-factor model in which all self-reported items were loaded onto a single latent factor. The results showed that the one-factor model exhibited poor fit, $\chi^2(299) = 2409.052$, CFI = 0.814, TLI = 0.798, RMSEA = 0.147, and SRMR = 0.098, and was substantially poorer than the proposed six-factor model. Taken together, the factor loadings, reliability coefficients, AVE values, and the comparison with the one-factor model provide support for the empirical distinctiveness of the constructs, reducing the concern that the observed relationships were solely driven by common method variance.

Table 4. Reliability and convergent validity of the constructs in Study 2.

Construct	Items	Standardized Factor Loadings	Cronbach's α	CR	AVE
Digital self-efficacy (DSE)	DSE1–DSE3	0.636–0.895	0.874	0.850	0.659
Cognitive load (CL)	CL1–CL6	0.632–0.894	0.971	0.891	0.580
Perceived control (PC)	PC1–PC4	0.900–0.936	0.794	0.956	0.845
Human–AI trust (TR)	TR1–TR5	0.884–0.945	0.967	0.958	0.820
Value co-creation (VCC)	VCC1–VCC4	0.893–0.945	0.912	0.960	0.858
Value co-destruction (VCD)	VCD1–VCD4	0.788–0.926	0.846	0.937	0.789

4.2.2. Moderation Analysis

To examine the moderating role of digital self-efficacy (DSE) in the two parallel paths, namely, the cognitive load path and the perceived control path, moderation analyses were conducted separately with cognitive load and perceived control as the dependent variables. The results are presented in Table 5.

Table 5. Moderating effects of digital self-efficacy on the dual-path mechanisms.

Path	GROUP $\times$ DSE	Conditional Effect	Bootstrap 95% CI
GROUP $\times$ DSE $\rightarrow$ CL	B = -0.555 , $p < 0.001$	Low DSE: -2.287^{***}	[-2.840 , -1.734]
		Medium DSE: -3.464^{***}	[-3.671 , -3.257]
		High DSE: -3.834^{***}	[-4.060 , -3.607]
GROUP $\times$ DSE $\rightarrow$ PC	B = 0.663 , $p < 0.001$	Low DSE: 1.509^{***}	[0.884 , 2.133]
		Medium DSE: 2.915^{***}	[2.681 , 3.148]
		High DSE: 3.357^{***}	[3.101 , 3.612]

Note: $*** p < 0.001$.

The results showed that the interaction term for the cognitive load path (GROUP $\times$ DSE $\rightarrow$ CL) was statistically significant (B = -0.555 , $p < 0.001$). Simple slope patterns indicate that the association between fit condition and cognitive load varied across levels of digital self-efficacy. Specifically, the negative association between fit condition and cognitive load showed increasing magnitude from low to high DSE (conditional association estimates increased from -2.287 to -3.834). The interaction term for the perceived control path (GROUP $\times$ DSE $\rightarrow$ PC) was also statistically significant (B = 0.663 , $p < 0.001$). The positive association between fit condition and perceived control also showed increasing magnitude across levels of DSE, with conditional association estimates increasing from 1.509 to 3.357 . These patterns are consistent with H6 and H7, respectively. To reduce concerns about multicollinearity, we computed variance inflation factors (VIFs) for all predictors in the first-stage models; all VIFs were below 2.5, indicating that the regression coefficient estimates were unlikely to be seriously affected by multicollinearity.

After examining the conditional associations between fit condition and the two first-stage psychological variables, the study further examined whether the proposed conditional indirect association patterns were observed in the full serial mediation framework.

4.2.3. Moderated Serial Mediation Analysis

We further tested the full moderated serial mediation model using PROCESS Model 83 (Table 6). Before including the mediators, the total association between GROUP and value co-creation was significantly positive, B = 2.659 , SE = 0.099 , $t = 26.807$, $p < 0.001$, 95% CI [2.464 , 2.854], $R^2 = 0.688$, $F(1, 326) = 718.613$, $p < 0.001$. The total association between GROUP and value co-destruction was significantly negative, B = -1.991 , SE = 0.110 , $t = -18.147$, $p < 0.001$, 95% CI [-2.207 , -1.775], $R^2 = 0.503$, $F(1, 326) = 329.321$, $p < 0.001$. After the serial mediation paths were included in the full models, the direct associations between GROUP and value co-creation and value co-destruction were no longer significant. This pattern is consistent with the proposed psychological process framework, suggesting that the association between fit condition and value outcomes was statistically explained to a considerable extent by cognitive load, perceived control, and human–AI trust. However, the observed patterns should be interpreted as conditional psychological associations rather than strict process confirmations. To improve the transparency and replicability of the moderated chain mediation analyses, the complete regression results are reported in the Supplementary Materials, including the regression equations for each mediator and dependent variable, full path coefficients, standard errors, t values, p -values, confidence

intervals, R^2 values, F statistics, direct effects, total effects, conditional indirect effects, and indices of moderated mediation. The results of the conditional indirect effects are presented in Table 6, and the complete regression results are provided in Supplementary Tables S1–S5.

Table 6. Conditional indirect effects of fit condition on value outcomes across different levels of digital self-efficacy.

Dependent Variable	Path	DSE Level	Conditional Indirect Effect	Bootstrap 95% CI	Index of Moderated Mediation [95% CI]
VCC	GROUP $\times$ DSE $\rightarrow$ CL $\rightarrow$ TR $\rightarrow$ VCC	Low (3.55)	0.751	[0.4599, 1.1283]	Index = 0.1822 [0.0931, 0.2760]
		Medium (5.67)	1.137	[0.8415, 1.5172]	
		High (6.33)	1.258	[0.9358, 1.6631]	
	GROUP $\times$ DSE $\rightarrow$ PC $\rightarrow$ TR $\rightarrow$ VCC	Low (3.55)	0.447	[0.2276, 0.7424]	Index = 0.1965 [0.0871, 0.3282]
		Medium (5.67)	0.864	[0.6040, 1.2094]	
		High (6.33)	0.995	[0.6895, 1.3960]	
VCD	GROUP $\times$ DSE $\rightarrow$ CL $\rightarrow$ TR $\rightarrow$ VCD	Low (3.55)	−0.433	[−0.6984, −0.2286]	Index = −0.1051 [−0.1732, −0.0471]
		Medium (5.67)	−0.656	[−0.9647, −0.3866]	
		High (6.33)	−0.726	[−1.0615, −0.4282]	
	GROUP $\times$ DSE $\rightarrow$ PC $\rightarrow$ TR $\rightarrow$ VCD	Low (3.55)	−0.223	[−0.4327, −0.0892]	Index = −0.0982 [−0.1816, −0.0381]
		Medium (5.67)	−0.432	[−0.7137, −0.2127]	
		High (6.33)	−0.497	[−0.8204, −0.2468]	

The complete regression results are consistent with the proposed moderated conditional indirect association pattern. For the cognitive load path, the interaction between GROUP and digital self-efficacy was statistically associated with cognitive load ($B = -0.555$, $SE = 0.113$, $t = -4.927$, $p < 0.001$), and cognitive load was negatively associated with human–AI trust ($B = -0.438$, $SE = 0.044$, $t = -9.937$, $p < 0.001$). Human–AI trust was positively associated with value co-creation ($B = 0.749$, $SE = 0.059$, $t = 12.628$, $p < 0.001$) and negatively associated with value co-destruction ($B = -0.432$, $SE = 0.078$, $t = -5.521$, $p < 0.001$). For the perceived control path, the interaction between GROUP and digital self-efficacy was statistically associated with perceived control ($B = 0.663$, $SE = 0.127$, $t = 5.218$, $p < 0.001$), and perceived control was positively associated with human–AI trust ($B = 0.388$, $SE = 0.039$, $t = 10.020$, $p < 0.001$). Human–AI trust was positively associated with value co-creation ($B = 0.763$, $SE = 0.060$, $t = 12.799$, $p < 0.001$) and negatively associated with value co-destruction ($B = -0.382$, $SE = 0.077$, $t = -4.943$, $p < 0.001$).

For value co-creation, the conditional indirect associations of both the cognitive load chain and the perceived control chain increased with digital self-efficacy (DSE) (Table 6). The indices for the two serial paths were significant (0.1822 and 0.1965, respectively, with 95% CIs excluding zero), indicating conditional differences in the indirect associations across levels of digital self-efficacy. Specifically, higher DSE was associated with stronger first-stage associations between fit condition and lower cognitive load, as well as between fit condition and higher perceived control. These conditional first-stage associations were further linked to higher value co-creation through human–AI trust.

For value co-destruction, the conditional indirect associations of both serial paths were negative, and their absolute values increased across levels of DSE. The indices for the two serial paths were also significant (−0.1051 and −0.0982, respectively, with 95% CIs excluding zero), indicating conditional differences in the indirect associations across levels of digital self-efficacy. Specifically, higher DSE was associated with stronger first-stage associations between fit condition and lower cognitive load, as well as between fit condition and higher perceived control. These conditional first-stage associations were further linked to lower value co-destruction through human–AI trust.

These findings reflect conditional psychological associations derived from experimental scenarios rather than temporal or causal process confirmation. Thus, H8 and H9 were statistically consistent with the observed conditional indirect association patterns for both value co-creation and value co-destruction. Accordingly, the moderated serial mediation model should be understood as a statistical representation of theoretically specified conditional associations, not as evidence that these psychological states unfold in a fixed temporal order.

4.2.4. Outcome Distinctiveness and Mechanism Interaction Checks

To further clarify the empirical relationship between value co-creation and value co-destruction, we conducted additional outcome-distinctiveness checks in Study 2. The Pearson correlation analysis showed that value co-creation and value co-destruction were strongly and negatively correlated, $r = -0.816$, $p < 0.001$. This high negative correlation indicates that the two outcomes were closely related in opposite directions in the scenario-based service context. Therefore, we do not interpret value co-creation and value co-destruction as fully independent or empirically distant constructs.

At the same time, the correlation does not necessarily imply that the two outcomes are perfectly reversible or measurement equivalent. First, the correlation was below 1.00, indicating that the two constructs did not completely overlap. Second, the square roots of AVE for value co-creation and value co-destruction were 0.926 and 0.888, respectively, both of which exceeded the absolute value of their correlation. Third, the six-factor measurement model showed better empirical distinctiveness than the one-factor model reported in Section 4.2.1. Taken together, these results suggest that value co-creation and value co-destruction should be interpreted as highly related but conceptually distinguishable outcomes. Accordingly, we retained them as two separate dependent variables, while explicitly acknowledging their strong negative association.

We also examined whether cognitive load and perceived control exhibited interaction dynamics. In Study 1, the interaction between cognitive load and perceived control did not significantly predict human–AI trust, $B = 0.011$, $t = 0.577$, $p = 0.564$, value co-creation, $B = 0.032$, $t = 1.463$, $p = 0.145$, or value co-destruction, $B = 0.005$, $t = 0.211$, $p = 0.833$. These results suggest that cognitive load and perceived control mainly operated as parallel psychological mechanisms in Study 1. In Study 2, the interaction between cognitive load and perceived control significantly predicted human–AI trust, $B = 0.037$, $t = 2.298$, $p = 0.022$, suggesting that perceived control may partially buffer the negative association between cognitive load and human–AI trust. However, after human–AI trust was included in the model, the interaction term did not significantly predict value co-creation, $B = -0.021$, $t = -1.140$, $p = 0.255$, nor did it reach the conventional 0.05 significance level in predicting value co-destruction, $B = 0.042$, $t = 1.776$, $p = 0.077$. Overall, these findings suggest that cognitive load and perceived control generally function as parallel psychological mechanisms, while Study 2 provides limited evidence that they may also interact at the stage of human–AI trust formation. The detailed results of these additional analyses are reported in Supplementary Table S6.

4.2.5. Robustness Checks, Exploratory Analyses, and Causal-Inference Boundaries

Before conducting the main analyses, we first examined whether participants in the high-fit and low-fit conditions differed in several pre-measured characteristics. Because the original occupation categories contained several small cells, occupation was recoded into enterprise employee status for the randomization check and robustness analyses. As shown in Table 7, the randomization check showed no significant differences between the

two conditions in gender, age, enterprise employee status, prior usage frequency, or digital self-efficacy. These results suggest that the random assignment was generally successful.

Table 7. Randomization checks and robustness analyses.

Analysis	Variable/Path	Estimate/Test Statistic	p-Value/95% CI
Randomization check	Gender	$\chi^2(1) = 1.578$	$p = 0.209$
	Age	$\chi^2(3) = 3.490$	$p = 0.322$
	Enterprise employee status	$\chi^2(1) = 0.118$	$p = 0.731$
	Prior usage frequency	$\chi^2(2) = 5.212$	$p = 0.074$
	Digital self-efficacy	$t = -1.283$	$p = 0.188$
Robustness analysis	Fit condition → Value co-creation	$B = 2.629$, $SE = 0.103$, $t = 25.459$	$p < 0.001$
	Fit condition → Value co-destruction	$B = -1.913$, $SE = 0.113$, $t = -16.899$	$p < 0.001$
	CL → TR → VCC	Index = 0.183	95% CI [0.090, 0.277]
	CL → TR → VCD	Index = -0.105	95% CI [-0.173, -0.045]
	PC → TR → VCC	Index = 0.199	95% CI [0.084, 0.328]
	PC → TR → VCD	Index = -0.101	95% CI [-0.188, -0.038]

Note. CL = cognitive load; PC = perceived control; TR = human–AI trust; VCC = value co-creation; VCD = value co-destruction. Control variables included gender, age, enterprise employee status, and prior usage frequency.

As a robustness check, we further re-estimated the main-effect and moderated serial mediation models after including gender, age, enterprise employee status, and prior usage frequency as control variables. Digital self-efficacy was retained as the focal moderator rather than treated as a demographic control. As shown in Table 7, the main effects of fit condition remained significant after adding these controls. The moderated serial mediation results also remained substantively unchanged, with all indices of moderated mediation remaining significant. These results indicate that the main findings were robust to demographic and usage-related controls.

We further compared the hypothesized structural model with several theoretically plausible alternative specifications. The hypothesized model specified cognitive load and perceived control as two parallel psychological mechanisms that influence value outcomes through human–AI trust. The alternative models included a parallel mediation model, in which cognitive load, perceived control, and human–AI trust were specified as parallel mediators; a no-trust model, in which cognitive load and perceived control directly predicted value outcomes without the trust conversion mechanism; a trust-first model, in which human–AI trust preceded cognitive load and perceived control; and a direct-only model, in which fit condition directly predicted value co-creation and value co-destruction without the proposed psychological mechanisms. The results are reported in Supplementary Table S7. Compared with the hypothesized model, the no-trust model and the direct-only model showed substantially poorer fit, suggesting that models excluding the proposed psychological mechanisms or the trust conversion mechanism were less consistent with the observed data. Although the parallel mediation model and the trust-first model showed comparable fit indices, the parallel mediation model produced an inadmissible solution, and the trust-first model produced a near-boundary solution. Therefore, these alternative specifications did not provide a more stable explanation than the hypothesized model.

We also conducted exploratory analyses to examine whether digital self-efficacy showed nonlinear or asymmetric patterns under the low-fit condition. As shown in Table 8, the tail-group analysis did not provide strong evidence for a nonlinear dark-side effect. Cognitive load showed a marginal increasing tendency across DSE groups, Welch $F(2, 63.823) = 2.983$, $p = 0.058$, whereas perceived control, human–AI trust, and value co-destruction did not differ significantly across groups. In particular, value co-destruction was not higher among high-DSE users under low-fit conditions. Thus, the exploratory results suggest that high-

DSE users may be more cognitively sensitive to misfit, but this sensitivity did not translate into significantly stronger value co-destruction.

Table 8. Exploratory analyses of digital self-efficacy.

Analysis	DV	Estimate/Test Statistic	Conditional Effects/J–N Result
TGA under low-fit condition	CL	Welch $F(2, 63.823) = 2.983$, $p = 0.058$	—
TGA under low-fit condition	PC	Not significant	—
TGA under low-fit condition	TR	Not significant	—
TGA under low-fit condition	VCD	Not significant	—
SS + J–N analysis	CL	GROUP $\times$ DSE: $B = -0.555$, SE = 0.113, $t = -4.927$, $p < 0.001$, 95% CI $[-0.777, -0.334]$	Low: $B = -2.526$, SE = 0.237, $p < 0.001$; Mean: $B = -3.270$, SE = 0.120, $p < 0.001$; High: $B = -4.013$, SE = 0.136, $p < 0.001$; J–N: no transition point within the observed DSE range
SS + J–N analysis	PC	GROUP $\times$ DSE: $B = 0.663$, SE = 0.127, $t = 5.218$, $p < 0.001$, 95% CI $[0.413, 0.913]$	Low: $B = 1.794$, SE = 0.267, $p < 0.001$; Mean: $B = 2.683$, SE = 0.136, $p < 0.001$; High: $B = 3.571$, SE = 0.153, $p < 0.001$; J–N: significant when DSE_A > 2.565, covering 94.51% of the sample

Note. TGA = tail-group analysis; SS = simple-slope analysis; J–N = Johnson–Neyman analysis; DV = dependent variable; DSE = digital self-efficacy; CL = cognitive load; PC = perceived control; TR = human–AI trust; VCD = value co-destruction. GROUP was coded as 0 = low fit and 1 = high fit.

To further probe this sensitivity pattern, we conducted simple-slope and Johnson–Neyman analyses for the two first-stage paths in which digital self-efficacy served as the moderator. As shown in Table 8, the interaction between fit condition and digital self-efficacy was significant for both cognitive load and perceived control. For cognitive load, the negative association between fit condition and cognitive load became stronger as DSE increased. The Johnson–Neyman analysis identified no statistical significance transition point within the observed range of DSE, indicating that the effect of fit condition on cognitive load remained significant across the full observed range of the moderator. For perceived control, the positive association between fit condition and perceived control also became stronger as DSE increased. The Johnson–Neyman analysis showed that the effect of fit condition became statistically significant when DSE_A exceeded 2.565, covering 94.51% of the sample. Overall, these exploratory analyses suggest that digital self-efficacy amplifies users' psychological sensitivity to fit and misfit. However, these results should be interpreted as preliminary evidence of possible asymmetry rather than definitive proof of a tipping point.

Finally, the causal-inference boundaries of the proposed model should be clarified. Perceived fit was operationalized through randomly assigned high-fit and low-fit conditions rather than through a post-scenario self-reported predictor. Therefore, the experimental design supports the interpretation that fit condition was associated with participants' post-scenario psychological evaluations and behavioral intentions. However, cognitive load, perceived control, human–AI trust, and value outcomes were all measured after scenario exposure; therefore, their serial ordering is theoretically specified rather than temporally established. Following Rohrer et al. [47], alternative model comparisons should be viewed only as robustness checks, because reversing arrows or comparing model fit indices cannot determine which causal model is true. Accordingly, the proposed model should be interpreted as a theoretically grounded and empirically supported process model, while stronger claims about mediator ordering require future longitudinal, field-experimental, or

interaction-log evidence. Therefore, the findings should be interpreted within a bounded framework. The present design allows us to examine whether theoretically specified psychological association patterns emerge under controlled scenario conditions, but it does not establish how these psychological states unfold moment by moment during actual service interactions. The framework is therefore intended to clarify possible conditional association patterns among perceived fit, cognitive load, perceived control, human–AI trust, and value outcomes, rather than to offer a universal causal account of intelligent customer service interactions.

4.3. Discussion of Study 2

Study 2 provided support for the moderating role of digital self-efficacy. The results showed that, compared with users with lower digital self-efficacy, users with higher digital self-efficacy exhibited stronger associations between fit condition and lower cognitive load, as well as between fit condition and higher perceived control. Along the cognitive load path, higher DSE appeared to strengthen the link between fit condition and reduced cognitive burden, which was further associated with human–AI trust and value outcomes. Along the perceived control path, users with higher DSE were more likely to interpret system task support as controllable assistance rather than as a loss of control. Thus, fit condition was more strongly associated with perceived control among high-DSE users, and this association was further linked to value co-creation and value co-destruction through human–AI trust. Overall, Study 2 suggests that digital self-efficacy functions as an important boundary condition in the proposed process model.

5. Discussion

Taken together, the findings clarify the originality of this study. The contribution of this research lies not only in examining both value co-creation and value co-destruction, but also in showing that these two outcomes are not explained by a fully symmetric psychological process. At the same time, the findings should not be interpreted as showing a clean or deterministic transition from one outcome to the other; rather, they indicate a probabilistic value divergence process shaped by cognitive burden, perceived agency, and trust formation. In intelligent customer service encounters, value co-creation is more likely to arise when perceived fit reduces interaction barriers, supports users' perceived control, and strengthens human–AI trust. By contrast, value co-destruction may be more likely when users report higher cognitive effort, weaker perceived control, lower trust, and stronger negative responses such as frustration, avoidance, or resistance. This suggests that value co-destruction should be regarded as an independent value trajectory rather than as a simple failure to achieve value co-creation. By integrating the processing-cost logic of cognitive load, the agency logic of perceived control, and the evaluative logic of human–AI trust, this study provides a more nuanced explanation of how the same intelligent customer service interface may lead to divergent value outcomes.

This asymmetry is important for interpreting the model. The results do not suggest that cognitive load and perceived control operate as two equally supported paths for both value outcomes. In Study 1, H2a was not supported: cognitive load did not directly explain value co-creation. Instead, cognitive load was more clearly associated with value co-destruction and operated in the value co-creation path mainly when further linked to human–AI trust. Perceived control, by contrast, showed a broader association with both value co-creation and value co-destruction. Thus, the dual-path model should be understood as an asymmetric process model rather than as a balanced set of parallel pathways.

The findings further suggest that value divergence should be interpreted as a bounded and contingent process rather than as a deterministic transition from value co-creation to value co-destruction. This study does not identify a precise psychological threshold at which intelligent customer service shifts from facilitating value co-creation to generating value co-destruction. However, the asymmetric pattern of the results indicates that such a transition may become more likely when cognitive burden accumulates, perceived control weakens, and trust in the system becomes unstable. In this sense, value co-destruction may emerge not simply because value co-creation is absent, but because users begin to interpret the automated service encounter as effortful, uncontrollable, or unreliable. At the same time, low perceived fit does not necessarily lead to value co-destruction in all cases. Users with stronger digital self-efficacy, access to human assistance, or sufficient coping resources may tolerate temporary misfit and prevent negative responses from escalating. Conversely, even a high-fit system may create perceived loss of control when automation becomes opaque, difficult to interrupt, or unsuitable for emotionally sensitive service tasks. Therefore, the proposed framework should be understood as identifying conditional and theoretically specified association patterns, rather than establishing a universal or temporally fixed causal sequence from perceived fit to value outcomes.

5.1. Theoretical Contributions

First, this study contributes to extending task–technology fit theory to the post-use value formation process of intelligent customer service. Prior research has mainly focused on users’ adoption of intelligent customer service or the direct effects of service characteristics on satisfaction and continuance intention. In contrast, this study uses perceived fit as the primary theoretical anchor to explain why the same intelligent customer service condition may be associated with divergent value outcomes during actual use. By linking fit condition to both value co-creation and value co-destruction, this study provides a more nuanced fit-based perspective for understanding user value divergence in human–AI service interactions.

Second, this study clarifies the asymmetric psychological mechanisms through which perceived fit is linked to value outcomes. Specifically, cognitive load and perceived control are introduced as two localized mechanisms within the task–technology fit framework. The findings suggest that these two mechanisms do not operate as fully symmetric or equally supported pathways. Cognitive load was more clearly associated with value co-destruction and was linked to value co-creation mainly through human–AI trust, whereas perceived control showed a broader association with both value co-creation and value co-destruction. This finding helps explain why value co-destruction should not be treated simply as the absence or reverse of value co-creation.

Third, this study identifies human–AI trust as a conversion mechanism in the value formation process. By examining the theoretically specified path of “perceived fit → cognitive load/perceived control → human–AI trust → value co-creation/value co-destruction”, this study shows how users’ interaction experiences are associated with trust judgments and subsequent value-related responses. Rather than treating trust as a uniformly positive mediator, the revised framework emphasizes that trust formation depends on how users interpret the balance between automation efficiency and perceived agency.

Fourth, this study highlights digital self-efficacy as an important boundary condition in intelligent customer service encounters. The findings suggest that users with higher digital self-efficacy show stronger associations between fit condition and lower cognitive load, stronger perceived control, higher human–AI trust, and more favorable value outcomes. This boundary condition helps explain why users may respond differently to the same

intelligent customer service condition and extends the application of self-efficacy theory to human–AI service contexts.

5.2. Managerial Implications

Given that this study is based on written scenario experiments and self-reported intentions, the following managerial implications should be interpreted as tentative design suggestions rather than validated decision-management prescriptions. They indicate possible directions for intelligent customer service design that require further validation using field experiments, behavioral data, interaction logs, and customer-service performance records. Nevertheless, the findings suggest that intelligent customer service should not be evaluated only through the lens of efficiency, automation rate, or cost reduction. A service system that is efficient for routine users may still exclude, frustrate, or disadvantage users with lower digital confidence, complex service requests, or emotionally sensitive complaints. Therefore, firms should be cautious that efficiency-oriented automation does not unintentionally create digitally exclusionary service experiences for vulnerable users.

First, firms should regard perceived fit in intelligent customer service as a core lever for improving human–AI interaction. Perceived fit is not merely a technical indicator; it is related to users' cognitive burden, perceived control, and trust formation, thereby influencing whether intelligent customer service is associated with value co-creation or value co-destruction. Therefore, firms should not only improve technical performance, such as natural language understanding and response accuracy, but also ensure that system responses align with users' task requirements, expectations, and mental models. An effective intelligent service system should both understand users' needs and remain manageable during the interaction.

Second, intelligent customer service design should be differentiated according to the nature and risk level of the service task. For transactional tasks, such as order tracking, invoice requests, address changes, and basic information searches, the priority should be to minimize cognitive load through short interaction paths, clear prompts, pre-filled information, and visible progress indicators. In contrast, for complex or relational tasks, such as complaints, refund disputes, service recovery, high-value after-sales issues, and emotionally sensitive inquiries, efficiency should not be achieved by forcing users to remain within automated loops. Firms should provide options such as revising requests, obtaining explanations, comparing alternative solutions, pausing automated processes, withdrawing or correcting previous inputs, and switching to human agents. These design strategies can reduce the risk that intelligent customer service becomes opaque, frustrating, or uncontrollable.

Third, firms may consider developing empirically testable human-handoff rules rather than treating manual intervention as a last resort. Escalation to human agents could be considered when users repeatedly express dissatisfaction, when the system misunderstands user intent across multiple turns, when users repeatedly rephrase the same request, when negative sentiment becomes salient, when the dispute involves a relatively high monetary value, or when the user repeatedly requests human assistance but fails to obtain it. In these situations, continued automation may increase cognitive load, weaken perceived control, and damage trust. A timely transfer to a human agent may therefore function as an important component of service recovery and consumer protection.

Fourth, firms should monitor early warning signals associated with value co-destruction. Although this study does not estimate a precise operational threshold, the findings suggest that firms should track behavioral and linguistic signals such as repeated input, repeated correction, long interaction loops, failed intent recognition, negative wording, abrupt exit, complaint initiation, repeated requests for human agents, and failed

handoff attempts. When these signals appear, the system may simplify the task, provide clearer guidance, restore users' sense of control, or offer immediate human support before the poor interaction is associated with stronger value co-destruction responses.

Finally, firms should adopt inclusive service strategies for users with different levels of digital self-efficacy. High-DSE users may benefit more from efficient automated services, whereas low-DSE users may require clearer step-by-step guidance, simpler language, visible progress cues, more controllable options, and easier access to human support. More broadly, intelligent customer service may unintentionally create unequal service experiences for users with low digital confidence, older consumers, users with limited digital literacy, and other marginalized groups. Therefore, firms should provide accessible interfaces, transparent explanations, multimodal channels, easy human handoff, and non-digital alternatives. In this sense, improving intelligent customer service is not only about efficiency or automation performance, but also about service fairness, consumer inclusion, and digital accessibility.

5.3. Limitations and Future Research

This study has several limitations. First, this study used online scenario-based experiments and self-reported intentions, which may differ from actual intelligent customer service encounters. Although this design allowed controlled manipulation of perceived fit, the scenarios represented relatively clear cases of successful versus unsuccessful service interactions, which may have made the contrast between conditions more salient than in natural settings. Thus, the large effect sizes should be interpreted mainly as evidence of experimentally induced differences and theoretically proposed psychological mechanisms, rather than as exact estimates of real-world effect magnitudes. Moreover, the experimental scenario was situated in a specific e-commerce service context, and the participants were consumers with prior experience using intelligent customer service. Therefore, the findings may not fully generalize to users with little or no prior intelligent customer service experience, to other service contexts such as healthcare, finance, public services, or offline service recovery, or to real-time interactions in which users' responses unfold dynamically. In addition, this study did not use actual interaction logs, behavioral records, customer-service performance data, or field experiments; therefore, the managerial implications should be further tested in real service environments before being treated as validated operational prescriptions. Future research could use more subtle scenario manipulations, include manipulation checks in each experimental wave, and combine field experiments, longitudinal designs, chatbot interaction logs, clickstream records, service recovery records, or repeated-use behavioral data to enhance ecological validity and examine how perceived fit, cognitive load, perceived control, human-AI trust, value co-creation, and value co-destruction evolve during actual service processes. To further delineate the limits of the proposed framework, this study should not be interpreted as establishing a fixed causal sequence from perceived fit to cognitive load/perceived control, then to human-AI trust, and finally to value outcomes. Rather, the framework identifies theoretically specified and empirically observed association patterns under scenario-based experimental conditions. Its explanatory scope is bounded by three conditions. First, the evidence is based on written service scenarios rather than real-time human-AI service interactions. Second, cognitive load, perceived control, human-AI trust, and value outcomes were measured through post-scenario self-reports rather than through temporally separated or behavioral measures. Third, the framework was developed in an e-commerce intelligent customer service context and may not fully generalize to more complex, high-risk, or emotionally sensitive service settings. Therefore, the proposed framework is best understood as a bounded and contingent explanation of value divergence, rather than as a universal causal model.

In addition, although the randomized scenario design strengthens the interpretation that fit condition shaped participants' post-scenario evaluations, the serial ordering among cognitive load, perceived control, human–AI trust, and value outcomes should be interpreted cautiously. These variables were measured through self-reported responses after exposure to the same scenario, which means that the data cannot fully rule out simultaneity, unobserved psychological factors, or common method influences among the mediators and outcomes. Therefore, the proposed model should be understood as a theoretically grounded process model supported by scenario-based evidence, rather than as definitive proof of a temporally ordered causal chain. Future research should use longitudinal designs, real-time chatbot interaction logs, field experiments, physiological or behavioral indicators of cognitive load, and repeated measures of trust and control to examine how these psychological states evolve during actual service failures and recoveries.

Second, regarding mechanism exploration, this study focuses on two paths, namely cognitive load and perceived control, but intelligent customer service may influence user value through other psychological mechanisms as well, such as affective experience and social presence. Future studies may further explore additional potential mediators and develop a more comprehensive theoretical model.

Third, in terms of the research sample and user heterogeneity, this study mainly focuses on consumers with prior experience using intelligent customer service, whereas different types of users may exhibit different behavioral patterns and psychological responses when interacting with intelligent customer service. In particular, generational differences may represent an important direction for future research. Users from different generational cohorts may differ in their familiarity with digital technologies, expectations of AI service agents, tolerance for interaction errors, and ability to manage cognitive demands during service encounters. Given that this study identifies cognitive load as a key psychological mechanism, future research could further examine whether users from different generations differ in their cognitive processing and responses to AI-enabled service interactions, and whether the effects of perceived fit on cognitive load, human–AI trust, value co-creation, and value co-destruction vary across generational groups.

Supplementary Materials: The following supporting information can be downloaded at: <https://www.mdpi.com/article/10.3390/jtaer21070225/s1>, Table S1: Total Effects of GROUP on Value Co-Creation and Value Co-Destruction in Study 2; Table S2: Full Regression Results for the Cognitive Load Chain Predicting Value Co-Creation in Study 2; Table S3: Full Regression Results for the Cognitive Load Chain Predicting Value Co-Destruction in Study 2; Table S4: Full Regression Results for the Perceived Control Chain Predicting Value Co-Creation in Study 2; Table S5: Full Regression Results for the Perceived Control Chain Predicting Value Co-Destruction in Study 2; Table S6: Additional Analyses of the Interaction Between Cognitive Load and Perceived Control; Table S7: Alternative Structural Model Comparisons.

Author Contributions: Conceptualization, J.R. and L.W.; Methodology, S.S.; Software, J.R.; Validation, J.R. and S.S.; Formal analysis, J.R.; Investigation, J.R.; Resources, Y.C.; Data curation, J.R.; Writing—original draft preparation, J.R.; Writing—review and editing, J.R.; Visualization, Y.C.; Supervision, J.R.; Project administration, J.R.; Funding acquisition, L.W. All authors have read and agreed to the published version of the manuscript.

Funding: This research was supported by the National Natural Science Foundation of China (NSFC) under the project “The Impact Effects and Value Recovery Mechanisms of Value Co-Destruction in Livestream E-commerce from the Consumer Perspective” (Grant No. 72462012), and by the Innovation Project of GUET Graduate Education under the project “Human–AI Collaboration or Conflict? The Influence Mechanism of Technology–Task Fit in Intelligent Customer Service on Users' Value Co-Creation and Co-Destruction” (Grant No. 2026YCXSI10).

Institutional Review Board Statement: This study was conducted in accordance with the Declaration of Helsinki and was approved by the Ethics Committee of Guilin University of Electronic Technology prior to data collection (approval/reference number: GUET202509110007). The study involved a scenario-based online experiment with adult users of intelligent customer service. All data were collected anonymously and used solely for academic research purposes.

Informed Consent Statement: Informed consent was obtained from all participants due to the use of anonymous questionnaires, and all participants remained unidentifiable.

Data Availability Statement: The datasets generated and/or analyzed during the current study are available from the corresponding author upon reasonable request, due to privacy and anonymity considerations.

Acknowledgments: The authors have no additional acknowledgments to declare.

Conflicts of Interest: The authors declare no conflicts of interest. The funder had no role in the design of the study; in the collection, analyses, or interpretation of data; in the writing of the manuscript; or in the decision to publish the results.

References

- Peng, G.C.; Wu, S.Y. Research on the mechanism of user avoidance behavior in the background of intelligent customer service: An exploration based on rooting theory and narrative research. *Inf. Stud. Theory Appl.* **2024**, *47*, 104–111. [\[CrossRef\]](#)
- Wu, J.F.; Zhu, Y.M.; Liu, Y.Y. Aversion effect for AI chatbot: Inducing factor, psychological mechanism and boundary conditions. *Nankai Bus. Rev.* **2023**, *26*, 179–191.
- Tuo, Y.Z.; Qin, B.B. Practical dilemmas and realization paths of artificial intelligence technology empowering tourists' well-being. *Tour. Trib.* **2023**, *38*, 3–6. [\[CrossRef\]](#)
- Singu, H.B.; Chakraborty, D.; Troise, C.; Camilleri, M.A.; Bresciani, S. Responsible AI for trustworthy tourism: A framework for mitigating ambiguity and anxiety with generative AI. *Technol. Forecast. Soc. Change* **2026**, *223*, 124407. [\[CrossRef\]](#)
- Kalai, A.T.; Nachum, O.; Vempala, S.S.; Zhang, E. Why language models hallucinate. *arXiv* **2025**, arXiv:2509.04664. [\[CrossRef\]](#)
- Huang, M.-H.; Rust, R.T. Artificial intelligence in service. *J. Serv. Res.* **2018**, *21*, 155–172. [\[CrossRef\]](#)
- George, S.R.; Manu, C.; Edward, M. Artificial intelligence in frontline service encounters: A systematic review and research agenda. *Int. J. Consum. Stud.* **2025**, *49*, e70048. [\[CrossRef\]](#)
- Coelho, E.C.d.S.; Farias, J.S. Value cocreation and codestruction in artificial intelligence-enabled service interactions: Literature review and research agenda. *Span. J. Mark. ESIC* **2025**, *29*, 444–464. [\[CrossRef\]](#)
- Gao, J.; Opute, A.P.; Jawad, C.; Zhan, M. The influence of artificial intelligence chatbot problem solving on customers' continued usage intention in e-commerce platforms: An expectation-confirmation model approach. *J. Bus. Res.* **2025**, *200*, 115661. [\[CrossRef\]](#)
- Pham, T.C.T.; Huynh, T.; Tran, P.-T.; Ly, A.Q. Exploring how AI chatbots influence customer loyalty among Generation Z on e-commerce platforms. *Adv. Consum. Res.* **2025**, *2*, 2473–2486.
- Kim, M.; Oh, J.; Kim, B. Experience of digital music services and digital self-efficacy among older adults: Enjoyment and anxiety as mediators. *Technol. Soc.* **2021**, *67*, 101773. [\[CrossRef\]](#)
- Plé, L.; Chumpitaz Cáceres, R. Not always co-creation: Introducing interactional co-destruction of value in service-dominant logic. *J. Serv. Mark.* **2010**, *24*, 430–437. [\[CrossRef\]](#)
- Schulz, T.; Zimmermann, S.; Böhm, M.; Gewald, H.; Krcmar, H. Value co-creation and co-destruction in service ecosystems: The case of the Reach Now app. *Technol. Forecast. Soc. Change* **2021**, *170*, 120926. [\[CrossRef\]](#)
- Wardani, S.M.; Balqiah, T.E.; Astuti, R.D. Customer-centric co-creation value for vulnerable populations in digital healthcare. *Health Mark. Q.* **2025**, *42*, 369–391. [\[CrossRef\]](#) [\[PubMed\]](#)
- Luo, J.; Wang, J.; Liu, R.; Li, K. Value co-destruction: The dark side of interactions in online brand communities. *SAGE Open* **2024**, *14*, 21582440241295808. [\[CrossRef\]](#)
- Alexander, M.; Vallström, N. Value co-destruction: Problems and solutions. *AMS Rev.* **2023**, *13*, 200–210. [\[CrossRef\]](#)
- Lumivalo, J.; Tuunanen, T.; Salo, M. Value co-destruction: A conceptual review and future research agenda. *J. Serv. Res.* **2024**, *27*, 159–176. [\[CrossRef\]](#)
- Wu, B.; Chen, X. Continuance intention to use MOOCs: Integrating the technology acceptance model (TAM) and task technology fit (TTF) model. *Comput. Hum. Behav.* **2017**, *67*, 221–232. [\[CrossRef\]](#)
- Spies, R.; Grobbelaar, S.; Botha, A. A scoping review of the application of the task-technology fit theory. In *Responsible Design, Implementation and Use of Information and Communication Technology*; Hattings, M., Matthee, M., Smuts, H., Pappas, I., Dwivedi, Y., Mäntymäki, M., Eds.; Lecture Notes in Computer Science; Springer: Cham, Switzerland, 2020; Volume 12066, pp. 397–408. [\[CrossRef\]](#)

20. Sonntag, M.; Mehmman, J.; Teuteberg, F. AI-based chatbots in customer service: A task-technology fit (TTF) model. *Int. J. Serv. Sci. Manag. Eng. Technol.* **2025**, *16*, 1–20. [[CrossRef](#)]
21. Lalicic, L.; Weismayer, C. Consumers' reasons and perceived value co-creation of using artificial intelligence-enabled travel service agents. *J. Bus. Res.* **2021**, *129*, 891–901. [[CrossRef](#)]
22. Li, B.; Liu, L.; Mao, W.; Qu, Y.; Chen, Y. Voice artificial intelligence service failure and customer complaint behavior: The mediation effect of customer emotion. *Electron. Commer. Res. Appl.* **2023**, *59*, 101261. [[CrossRef](#)]
23. Aminian, N.O.; Michael, W.; Ghoreyshi, M. Cognitive workload assessment: A comprehensive review. *J. Aircr. Spacecr. Technol.* **2025**, *9*, 12–33. [[CrossRef](#)]
24. Sweller, J. Cognitive load theory and individual differences. *Learn. Individ. Differ.* **2024**, *110*, 102423. [[CrossRef](#)]
25. Evans, P.; Vansteenkiste, M.; Parker, P.; Kingsford-Smith, A.; Zhou, S. Cognitive load theory and its relationships with motivation: A self-determination theory perspective. *Educ. Psychol. Rev.* **2024**, *36*, 7. [[CrossRef](#)]
26. Brachten, F.; Brünker, F.; Frick, N.R.; Ross, B.; Stieglitz, S. On the ability of virtual agents to decrease cognitive load: An experimental study. *Inf. Syst. E-Bus. Manag.* **2020**, *18*, 187–207. [[CrossRef](#)]
27. Chang, H.Y.; Wu, H.F.; Chang, Y.C.; Tseng, Y.S.; Wang, Y.C. The effects of a virtual simulation-based, mobile technology application on nursing students' learning achievement and cognitive load: Randomized controlled trial. *Int. J. Nurs. Stud.* **2021**, *120*, 103948. [[CrossRef](#)] [[PubMed](#)]
28. Legaspi, R.; Xu, W.; Konishi, T.; Wada, S.; Kobayashi, N.; Naruse, Y.; Ishikawa, Y. The sense of agency in human–AI interactions. *Knowl.-Based Syst.* **2024**, *286*, 111298. [[CrossRef](#)]
29. Ryan, R.M.; Deci, E.L. Self-determination theory and the facilitation of intrinsic motivation, social development, and well-being. *Am. Psychol.* **2000**, *55*, 68–78. [[CrossRef](#)] [[PubMed](#)]
30. Dogru, E.Ö.; Krämer, N.C. Investigating appropriate reliance on AI-based decision support systems: The role of expertise, trust, and self-confidence. *J. Decis. Syst.* **2025**, *34*. [[CrossRef](#)]
31. Han, J.; Ko, D. Trust formation, error impact, and repair in human–AI financial advisory: A dynamic behavioral analysis. *Behav. Sci.* **2025**, *15*, 1370. [[CrossRef](#)] [[PubMed](#)]
32. Jones-Jang, S.M.; Park, Y.J. How do people react to AI failure? Automation bias, algorithmic aversion, and perceived controllability. *J. Comput.-Mediat. Commun.* **2023**, *28*, zmac029. [[CrossRef](#)]
33. Romeo, G.; Conti, D. Exploring automation bias in human–AI collaboration: A review and implications for explainable AI. *AI Soc.* **2026**, *41*, 259–278. [[CrossRef](#)]
34. Lee, J.D.; See, K.A. Trust in automation: Designing for appropriate reliance. *Hum. Factors* **2004**, *46*, 50–80. [[CrossRef](#)] [[PubMed](#)]
35. Ding, Y.; Najaf, M. Interactivity, humanness, and trust: A psychological approach to AI chatbot adoption in e-commerce. *BMC Psychol.* **2024**, *12*, 595. [[CrossRef](#)] [[PubMed](#)]
36. Patton, C.E. Automation use: The roles of self-confidence, trust, and workload. *Proc. Hum. Factors Ergon. Soc. Annu. Meet.* **2023**, *67*, 965–970. [[CrossRef](#)]
37. Goodhue, D.L.; Thompson, R.L. Task-technology fit and individual performance. *MIS Q.* **1995**, *19*, 213–236. [[CrossRef](#)] [[PubMed](#)]
38. Cohen, J. *Statistical Power Analysis for the Behavioral Sciences*, 2nd ed.; Lawrence Erlbaum Associates: Hillsdale, NJ, USA, 1988.
39. Hong, W.; Thong, J.Y.L.; Tam, K.Y. The effects of information format and shopping task on consumers' online shopping behavior: A cognitive fit perspective. *J. Manag. Inf. Syst.* **2004**, *21*, 149–184. [[CrossRef](#)]
40. Collier, J.E.; Barnes, D.C. Self-service delight: Exploring the hedonic aspects of self-service. *J. Bus. Res.* **2015**, *68*, 986–993. [[CrossRef](#)]
41. Lubbe, I.; Roberts-Lombard, M.; Langerman, J. Millennials' experiences and satisfaction with chatbots: A study of self-service technology in emerging markets. *Eur. Bus. Rev.* **2025**, *37*, 741–769. [[CrossRef](#)]
42. Hu, X.; Wu, G.; Wu, Y.; Zhang, H. The effects of third-party Web assurance seals on consumers' initial trust in online vendors: A functional perspective. *Decis. Support Syst.* **2010**, *48*, 407–418. [[CrossRef](#)]
43. Alagarsamy, S.; Mehrolia, S. Exploring chatbot trust: Antecedents and behavioural outcomes. *Heliyon* **2023**, *9*, e16074. [[CrossRef](#)] [[PubMed](#)]
44. Yang, X.C.; Tu, K. Research on the effect of platform support quality on user value co-creation citizenship behavior under the background of sharing economy. *Bus. Manag. J.* **2018**, *40*, 128–144. [[CrossRef](#)]
45. Vargo, S.L.; Lusch, R.F. Service-dominant logic: Continuing the evolution. *J. Acad. Mark. Sci.* **2008**, *36*, 1–10. [[CrossRef](#)]
46. Hayes, A.F. *Introduction to Mediation, Moderation, and Conditional Process Analysis: A Regression-Based Approach*, 3rd ed.; Guilford Press: New York, NY, USA, 2022.
47. Rohrer, J.M.; Hünermund, P.; Arslan, R.C.; Elson, M. That's a Lot to Process! Pitfalls of Popular Path Models. *Adv. Methods Pract. Psychol. Sci.* **2022**, *5*, 25152459221095827. [[CrossRef](#)]

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.